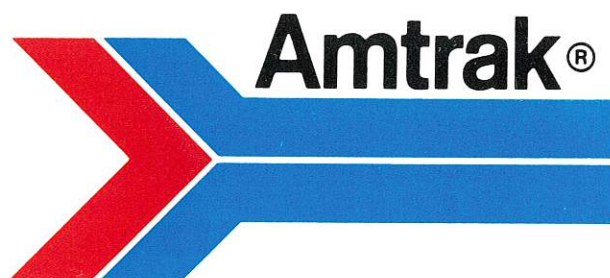


NATIONAL RAILROAD PASSENGER CORPORATION



THE ART OF SCHEDULING TRAINS

AND

THE IMPACT OF DELAYS

**National Railroad Passenger Corporation
Contract Management
March 27, 1996**

**© Copyright, National Railroad Passenger Corp. 1996
All Rights Reserved.**

Introduction and Acknowledgments

During 1995, at the request of Private Car Owners, I prepared a paper on the Art of Scheduling Trains. It was presented to the Private Car Owners at Nelson, British Columbia and to the Lexington Group in Transportation History at Portland, Oregon. The strong response we received prompted us to rewrite the paper as a more technical presentation on scheduling and the impact of delays.

Several years ago we developed a course on scheduling and schedule analysis that is still used on a regular basis for Amtrak management, and upon request it is also presented to the freight railroad personnel who are involved in analysis of Amtrak schedules.

This paper and a companion paper on the Art of Schedule Analysis have been prepared to provide direction to the course on scheduling and schedule analysis and to assist those who will be trained to conduct schedule evaluations.

I wish to acknowledge the assistance of Arthur Prentiss and Al Walton of my staff who carefully critiqued the draft. I also wish to acknowledge the assistance of Chris Dort and Rande Jones, locomotive engineers for Amtrak who reviewed the final draft. I also appreciated the editorial assistance from R. H. Young, Assistant Vice President Passenger Services, CSX; Howard Godwin, Director of Passenger Operations, CSX and Ray Shields, Manager of Passenger Operations, all of whom are certified locomotive engineers as well as Cliff Russ, Director of Passenger Services for CSX, their input was most valuable.

I appreciate permission from CSX to use the excerpts from their operating timetable.

James L. Larson
Washington, D.C.
March 27, 1996

Contents

Page

Part I: The Art of Scheduling

Background	1
Basic Principles	2
Pure Running Time Determination	6
Dwell Time	9
Recovery Time	12
Passenger Train Meets	13
Other Scheduling Considerations	14
Schedule Stringing	15
String Graphs	18
Connectivity	19
Negotiations	19

Part II: The Impact of Delays

The Impact of Delays	22
Signal Delays	24
Slow Order and Seasonal Schedules	25
Coordinating Maintenance Projects	26
Engineering Practices	27
Minimizing Delays from M of W Practices	28
Passenger Train Meets	31
Out of Slot	32
Heat Orders and Cold Orders	34
Avoidable Amtrak Related Delays	35
Killing Time	36
Late Connections	37
Selective Delay Reporting	38
Train Dispatching	40
Rationalization	42
Summary	43

PART I

THE ART OF SCHEDULING

Background:

In 1971, Amtrak inherited the schedules in place at that time on each of the carriers where passenger service was assumed. Some of the schedules were good, some were poor and some needed restringing. Thereafter, performance began to deteriorate. That deterioration continued until the first incentive contracts were negotiated with a limited number of carriers in 1974. However, no concerted effort was made to address the quality of schedules until the Second Amendment contracts were negotiated starting in 1976.

The carriers' solution to improved scheduling was to add more time to the schedules. This, in turn, led to more abuses of passenger train priority which, in turn, led to later trains. Since all performance was measured at endpoints, the carriers sought schedules that were extremely tight to the next to last station and then had virtually all recovery time built into the final terminal on that carrier. That led to performance where a train might be 45 minutes late at the next to last station yet arrive at the final station 20 minutes ahead of time.

The period from 1974 to 1976 under the first incentive agreements was an important learning period for both Amtrak and the freight carriers. There was a growing awareness that properly strung schedules and the elimination of excess time promoted improved performance. Since 1976, Amtrak's schedules have been negotiated with the carriers just as compensation provisions are negotiated.

Subsequently, the addition of intermediate point performance and performance penalties in 1981 further refined the process. When it comes to preparing proper schedules that both Amtrak and the railroads can accept, much has been learned over the years. What was unique about this process was the widespread operation of trackage rights, over 24,000 route miles, where passenger trains were operated over the infrastructure maintained and controlled (dispatched) by another company.

The relationship between Amtrak and the freight railroads has matured over the years and although the working relationship may vary significantly from one carrier to another, in general the working relationship has materially improved, as has the process of developing schedules.

Basic Principles:

The basic principles underlying proper scheduling are quite simple. Too much time in the schedule, the train will run late. Not enough time in the schedule, the train will run late. And with the right amount of time in the schedule, the train may still run late, depending upon delays. The key: (1) having the correct amount of time in the schedule and (2) having the schedule properly strung.

Even though Amtrak is frequently criticized in the media, the fact still remains that nine out of ten Metroliners arrive on time, eight out of ten Northeast Corridor trains arrive on time, and seven out of ten trains overall arrive on time.

Scheduling has three basic components:

- Pure Running Time
- Dwell Time
- Recovery Time

Pure Running Time is based upon running time developed at maximum authorized speeds, without delays, on a predetermined power to weight ratio. The power to weight ratio is the relationship of locomotive horsepower to trailing tonnage. It is a critical factor in scheduling trains.

Pure running time will vary depending upon the maximum authorized speeds and the power to weight ratio; therefore, these factors must be established before the pure running time can be determined.

Establishing the power-to-weight ratio has been somewhat oversimplified in the past by using the number of locomotive units to trailing cars, without regard to the units or the types of cars. Since the F-40 was the predominant locomotive type outside of electrified territory, it was fairly consistent in the establishment of power to weight ratios on all trains operated with F-40 locomotives. The more recent introduction of other diesel locomotive types such as the Dash-8 or P-40 locomotives, which have significantly different power and operating characteristics, has necessitated further

refinement to reflect the "normal" power consist operated on specific trains. Car weights vary generally between 50 tons and 74 tons, depending upon type, but the type of cars operated on a given train normally does not change from day-to-day. However, the number of cars often does.

Each train has a published "normal" consist, and while this establishes a good base for building a power to weight ratio, it should not be used as an average train. Additional cars may be required and, therefore, every time a car is added beyond the "normal" consist, it could adversely affect the power to weight ratio. It is good practice to allow for one or two additional cars in building a schedule so that extra cars can be operated without adversely affecting the performance of the train.

The additional time required for each additional car is not a straight line relationship. For example, with one F-40 locomotive and from four to six cars, the pure running time differences should be negligible. The seventh car may show some differential, depending upon the territory, while the eighth and ninth cars will definitely show a differential and the tenth car is not recommended.

It is interesting to note that the performance of one F-40 and six cars will not be materially different from two F-40's and twelve cars. Although the second train has six extra cars, about 500 feet in length, to pull through each restriction, only one unit provides head end power. Therefore the extra horsepower available for traction

provided by the second unit will generally compensate for the additional train length, as will a third unit with 18 cars.

Summary:

In developing a power to weight ratio, the type of power normally operated on the train, whether it is turbo, electric, or diesel should be used. This should be further refined by type, such as E-60 or AEM-7 or, for diesels, FL-9, F-40, Dash-8, or P-40, or any combination thereof. In selecting the number of trailing cars, the "normal" consist should be used, plus one or two cars to allow for extra cars in addition to the "normal" consist.

Pure Running Time Determination:

Once the power to weight ratio is determined, it can be applied to the authorized speeds to develop the pure running time, generally in one of three ways, as follows:

1. Field Riding Program
2. Manual Desk Calculation
3. Computerized Train Performance Calculation (TPC)

What is amazing is that regardless of the methodology used, the end results are usually very close -- generally within one or two percent.

The field riding program is performed by a team that rides the engine over the entire trip, factors out delays, and measures pure running times and dwell times to the nearest tenth of a minute. This is discussed in detail in a companion paper, "The Art of Schedule Analysis."

The manual desk calculation can be performed by only a very limited number of people who can manually calculate performance including acceleration, braking, curve speeds and other restrictions generally to the nearest tenth of a minute. If done correctly, it can be very accurate. This is used primarily to establish new service where unfinished capital work or other factors prevent actual field tests at prospective speeds.

Finally, TPC programs allow the computer to calculate performance over a line of railroad with amazing accuracy, provided accurate data are used in the program.

TPC programs are recommended and although they may or may not be as accurate as field tests, which measure actual performance, they are a great time-saver compared to multiple trips on long distance trains. They are especially useful in developing additional running time required for changing power to weight ratios. As an example, a TPC evaluation for a train with 2 units and 12 cars can be checked for accuracy with a field test with the same size consist and then used to determine the additional time required for 2 units and 14 cars, 16 cars and 18 cars. The user is cautioned that if the correct coefficients are not used, the results will not be accurate. The following three examples demonstrate common errors:

- (1) The original passenger car rolling resistance factor included a factor for the Spicer Drive mechanical transmission under each car which was used to generate power, using roughly 50 horsepower per car; however, modern equipment which uses head end power (HEP) does not have a Spicer Drive. Therefore, if the old rolling resistance factor is used in addition to factoring in the traction horsepower reduction for the HEP load, the resulting factor is a double count.

(2) Head end power is drawn off of only one unit, with a resulting reduction in horsepower available for traction. This train draws only as much power as is actually needed for HEP. The absolute maximum draw is about 625 horsepower from one unit. If horsepower needed to operate HEP is overcounted, the horsepower available for traction is understated.

(3) Some freight railroads use TPC programs which reflect normal freight train braking procedures, which, if applied to passenger trains, substantially overstate the pure running time.

Again, the accuracy of TPC programs should not be understated, provided errors such as those shown above are avoided. Amtrak once had the opportunity to review a TPC printout while making a field test between Seattle and Portland. It was surprising how close the two sets of data were. The TPC, however, failed to reflect that the station at Tacoma had been relocated some years earlier, so it is equally important to have up-to-date and accurate engineering data.

It should be remembered that TPC printouts only produce pure running time and do not include dwell time and recovery time. Recently, one carrier produced a printout that also included hypothetical en route delays which, of course, made the document worthless.

Dwell Time:

This is the amount of time required to perform station work. Smaller stations generally require three minutes or less, as does a crew change. Larger stations require more time depending upon train and locomotive servicing and/or switching, mail, baggage and express.

Dwell time is the second major component in developing satisfactory schedules. Dwell time is important to provide adequate time (but not excess time) to perform station work. It will vary significantly from one station to the next, depending upon requirements. In like manner, it may vary significantly at the same station from one day to the next. Some dwell times are only one minute at small stations. A crew change can be accomplished normally within three minutes. If the number varies significantly from one day to the next, one way to obtain a "normal" dwell is to average the times over a representative period. However this produces an "average," which means that half the time the train is on target or under, and half the time the train uses more than the allotted dwell time. Amtrak recommends the "75% Rule," which says: "Look at the time in the 75th percentile as the normal amount of dwell time required -- the balance is covered by recovery time."

There is a hypothesis that says: "Station work tends to expand to fill all available time." This says that if 15 minutes is currently required at Station "A" and 5 minutes additional time is added, making the scheduled dwell time 20 minutes, the work will, in a

short period of time, expand (or the pace of the work will slow down) to require 20 minutes -- and the recovery time is lost. Adding dwell time beyond actual requirements is not a good practice.

In like manner, adding excess dwell time because there are a lot of skiers on Fridays and Sundays during the season or because the train occasionally carries groups of boy scouts in the summertime is also not a good practice. This adds excess time 180 or 365 days a year for something that occurs perhaps 40 times a year. Again, this is the purpose of recovery time.

What about second stops at a station? Second stops are generally considered evil but may be a necessity. An extra stop generally costs less than two minutes, about one minute to pull forward and a minute or less for loading. Sometimes, the length of the platform or the makeup of a train necessitate second or even third spots. At unstaffed stations, the passengers tend to gather together under the platform light near the closed station building, and even the longest platform will not result in the passengers locating themselves along the platform. At those stations, it is doubtful that even a "wagon locator" such as is used in Europe would make any difference.

Second stops delay trains and therefore their use should be judicious and only as conditions require. At staffed stations, advance notice to the passengers on where to stand should be universal; often it is not done and thus additional stops are required. Alexandria, Virginia, is probably the best station in the nation for having passengers

properly spotted on the platform. Clearly, a second spot is better than walking passengers with their luggage through several cars on a moving train. Conversely, making a second, or even a third spot, just so a crew member can conveniently board the rear of the train is unacceptable.

The key is that if a second or third stop is made on a fairly regular basis, 50 percent or more, the time for the stop should be built into the dwell time. If it occurs seldom or only on a random basis, it should be covered by recovery time.

Other practices such as "streetcarring" passengers (collecting tickets on the platform) can have a disastrous impact upon dwell time. Finally, dwell time requirements may change over a period of time and must be revisited periodically. All employees must be monitored, encouraged and directed where necessary to minimize dwell times.

Summary:

Dwell time should be established to meet the exigencies of the service. Excess dwell time results in the station work expanding (or slowing down) to fill the available time. Additional spots are sometimes necessary, but their use should be judicious and incorporated into dwell times when necessary. Finally, dwell time requirements may change and scheduled dwell time must be adjusted on a recurring basis.

Recovery Time:

Recovery time is time added to the schedule, generally six percent¹ to eight percent, to compensate for all contingencies (delays) encountered en route. It may be as much as three hours per trip on a long distance train.

Recovery time by any other name -- fat, rubber, pad, etc. -- is still recovery time, and it is the third critical factor in proper schedule construction. It is an essential element, and how that time is spaced is addressed in the section on schedule stringing. The key point is:

WHEN THE AMOUNT OF DELAY EXCEEDS THE AVAILABLE RECOVERY TIME, THE TRAIN IS LATE.

Negative recovery time will be discussed further in the section on schedule stringing; however it can be defined as the difference in time when pure running time plus dwell time exceeds schedule time. To put this in perspective, if the pure running time between A and B is 40 minutes, and B has 5 minutes dwell time but the schedule is 43 minutes, there is -2.0 minutes recovery time or a negative recovery time of 2 minutes. From a practical standpoint, it means that with a perfect run with no delay and a stop equal to the scheduled dwell time, the train would be two minutes late. Negative

¹ Recovery time can be measured as a percent of pure running time only, or as a percent of pure running time plus dwell time, which together are sometimes called the "base schedule."

recovery time was an old scheduling practice to keep the schedule tight; the train would never wait for time because it could never make its schedule. It is not currently a recommended practice because it ensures the train being late. If a train needs 40 minutes running time and 5 minutes dwell time, then the schedule should be not less than 45 minutes.

Passenger Train Meets:

Beyond the three basic scheduling components -- pure running time, dwell time, and recovery time -- certain other allowances are necessary for proper scheduling. This includes passenger train meets. When such meets occur on single track, time is factored into the schedule to accommodate the meet. Which train will take siding will vary from day to day; and if the trains have equal priority, not even a crystal ball will predict a day in advance where the meet will occur or which train will take siding. Since 76 percent of the Amtrak system operations are in traffic control territory (TCS), the practice is to add 5 minutes additional recovery time to each train for the meet. While this becomes a windfall for the train holding the main track, provided it has clear signals, it mitigates the impact upon the train taking siding; the balance of the delay, if any, is covered by recovery time. Amtrak may request that some trains, especially those running toward critical connections with other passenger trains, be given priority in meets with their counterparts in the opposite direction.

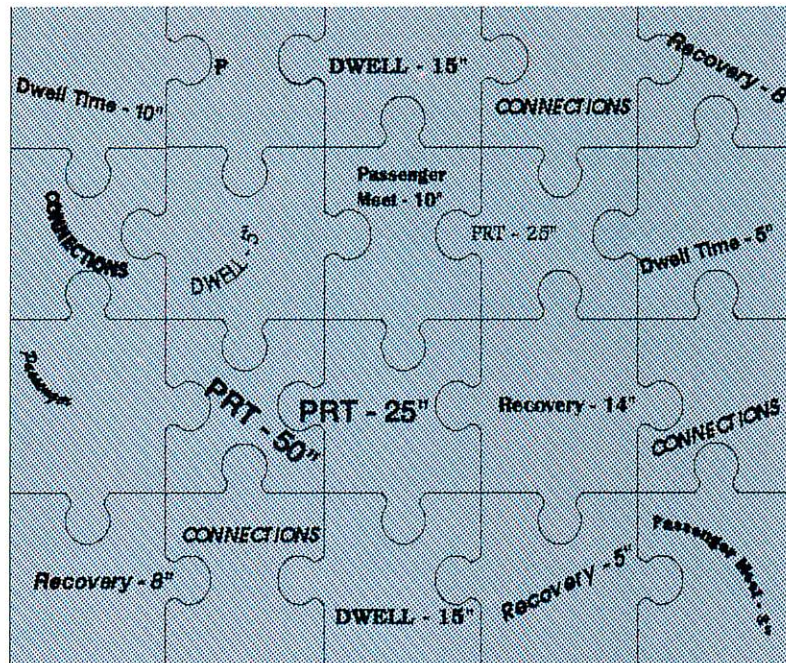
Other Scheduling Considerations:

Each train may encounter certain conditions peculiar to that route or unique to that train which may require special consideration in building a schedule. One example is Customs and Immigration delay, which, technically, is dwell time, and which occurs at different stations depending upon the direction of the train.

Probably, the most unusual example is the inclusion of the 5 extra minutes in the schedule of Train 3 to stop and cool traction motors at the summit of Raton Pass in Colorado (one direction only), as the ascending grade at 3-1/2 percent is the steepest grade on the Amtrak system.

Examples such as the ones cited above are rare; therefore, this discussion on schedule construction will be limited to pure running time, dwell time and recovery time, plus passenger train meets.

Schedule Stringing:



Now that the pure running time and dwell time requirements have been identified and the appropriate amount of recovery time determined, what will be done with it? The analogy of putting together a jigsaw puzzle is very close to the task of developing a schedule.

The critical skill is stringing the schedule. This means putting the time together in a meaningful way that permits the train to operate, generally without “killing time”² and yet reasonably close to one time. This is “the art of scheduling.”

Proper stringing of the schedule is critical because it plays a significant role in whether the train operates on time or late. It is also important to maintain on-time

² Waiting for scheduled departure time at a station or running slow in advance of a station.

performance at en route stations. When all the recovery time is bunched at the end point, the train becomes later and later en route as delays occur, but it arrives at its end point on-time or ahead of time. This is a disservice to all passengers entraining or detraining en route, which is often the overwhelming preponderance of the passengers, but it does help performance statistics and is the preferred methodology of freight railroads. In one extreme case on a long-distance train, the train could be 45 minutes late at the next to last station and ahead of time at its final station. This has since been corrected, but must be continually watched to avoid a recurrence.

Pure running times and dwell times are essentially a given, so the real variable is recovery time and how it is distributed. This discussion previously touched upon negative recovery time and took the position that it is not a good practice because it guarantees a late train in order to "preserve" recovery time for later use. Clearly, it is desirable to keep the schedule times "tight" (but not negative) at the first and second stations on a segment to avoid killing time which may be needed later. A well-strung schedule will be scheduled tight at the first one or two stations but will then distribute one-third to one-half of the recovery time before the intermediate contract railroad incentive checkpoint(s) or the end point. There is no ironclad rule to stringing recovery time; good judgment must be exercised.

What about killing time, which often is perceived as being evil. It is not bad in moderation, for example, one or two minutes. It only means the train is on-time. However, it would not be desirable to have a schedule where the train would kill time at

every station. That would be a bad schedule, and inordinately long. Killing excess amounts of time, say five or ten minutes or more on a regular basis, is also a bad practice and should be corrected through restringing, except at key stations where switching or major servicing is performed. Remember, a well-strung schedule materially contributes to a good operation.

As previously mentioned, an allowance should be made for passenger train meets that occur on single track. That allowance, which is actually additional recovery time, should be built into the schedule one or even two stations beyond the scheduled meeting point in the event that if one or the other train is late and the meet does not occur at the intended siding, the recovery time is not wasted by killing time prior to where the meet actually takes place.

Summary:

"WHEN THE AMOUNT OF DELAY EXCEEDS THE AVAILABLE RECOVERY TIME, THE TRAIN IS LATE." A well strung schedule will distribute the recovery time judiciously and will avoid extreme bunching of the recovery at the endpoint.

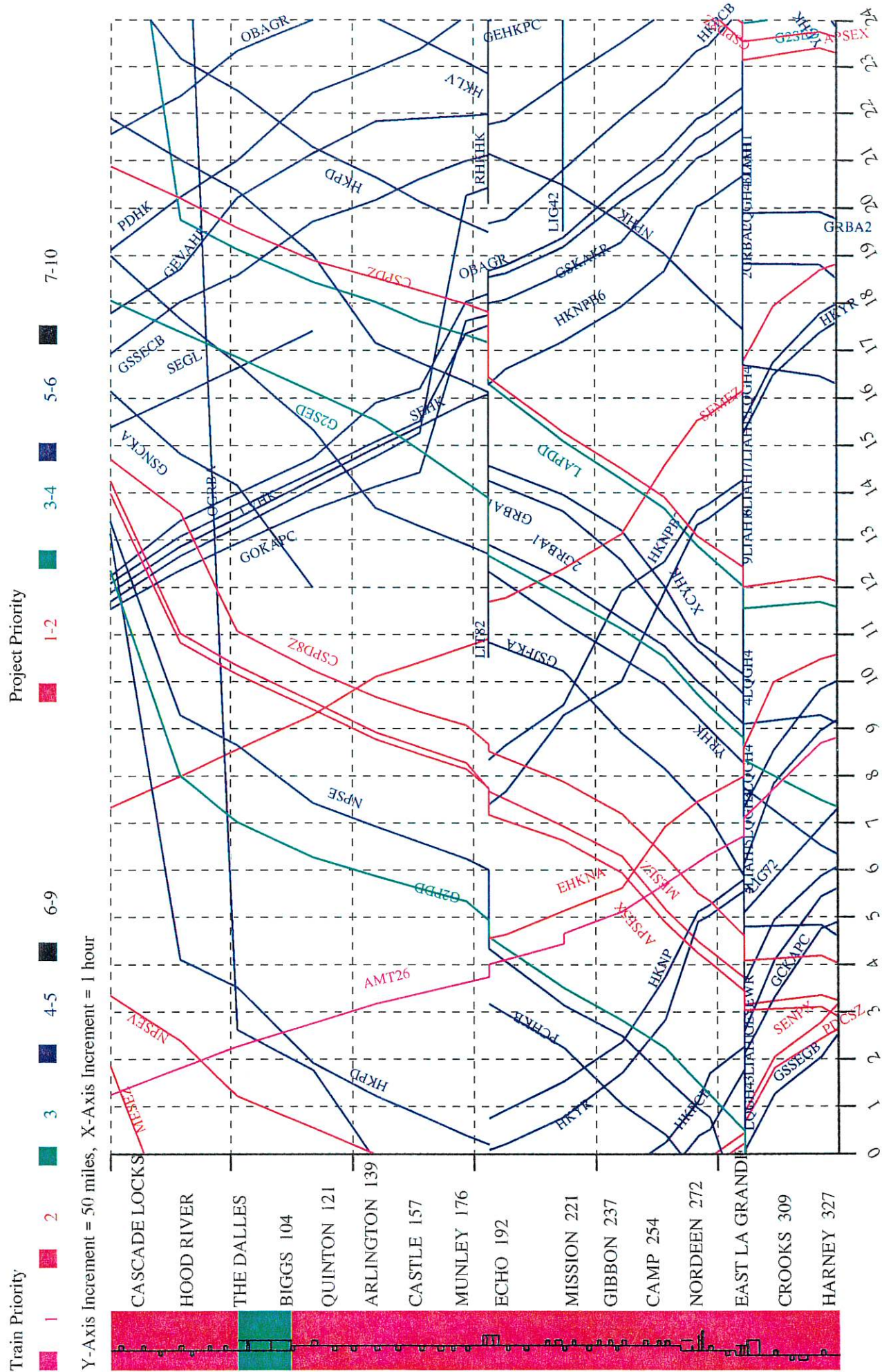
String Graphs:

Once schedules for a single-track railroad have been strung, it is a good practice to prepare a string graph. A string graph has time on one axis and stations or mileposts on the other axis. It is a graph of the operation. Some countries use string graphs as their official operating timetables.

The term "string graph" comes from a large panel used over 100 years ago to graph the trains with yarn. The original intent of a string graph was to avoid "cornfield meets" (head-on collisions). Today, it is used to determine if there is a siding available to facilitate meets, to identify passing locations, and to identify and avoid congestion points. When lines cross on a string graph it is a meet between opposing trains or a faster train overtaking a slower train.

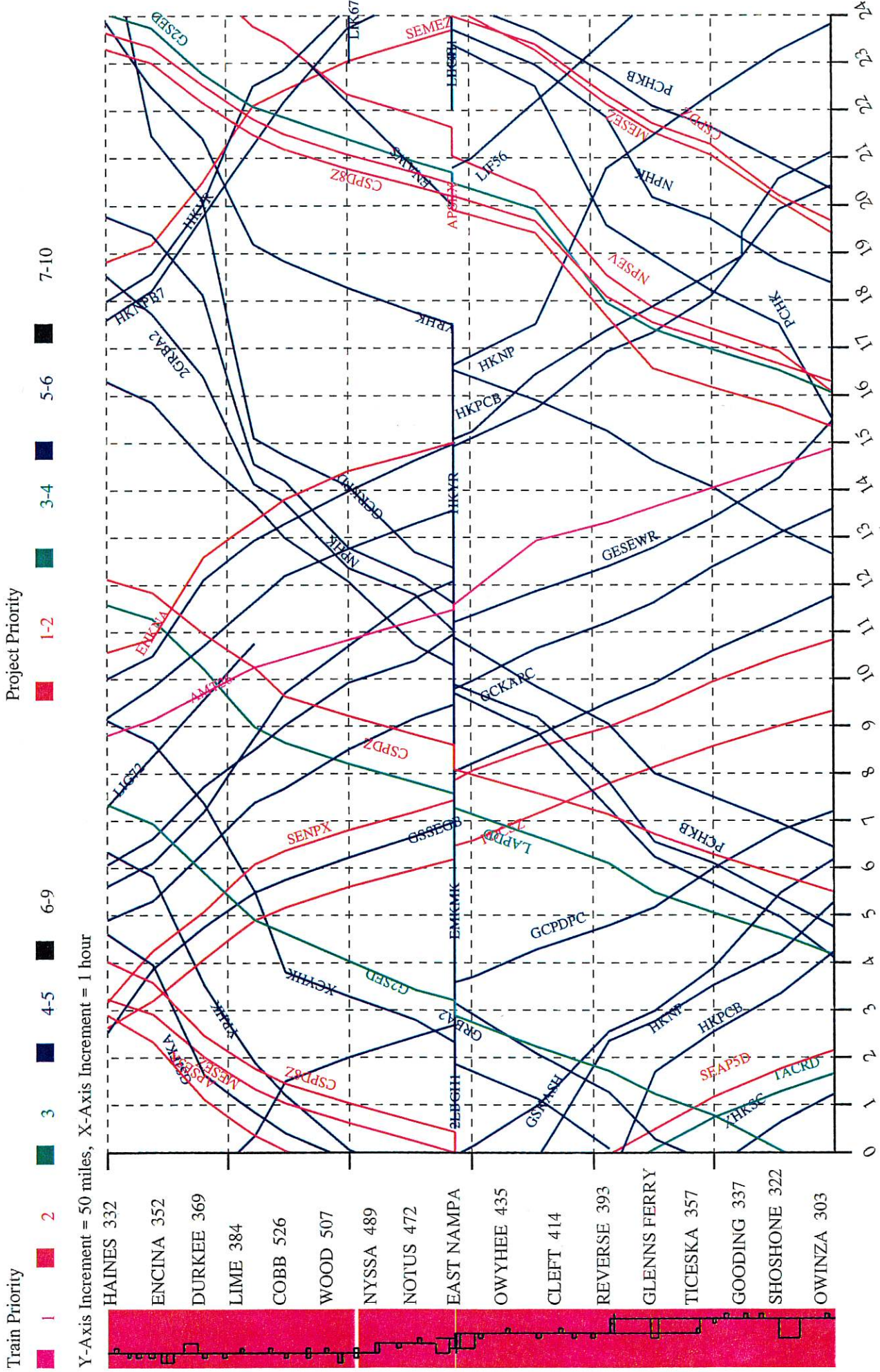
With modern day technology, computers are capable of preparing complicated string graphs. This procedure is currently used throughout the industry. There are a number of different systems in use. Amtrak uses PROTIM, which is a series of programs and databases developed by British Rail Business Systems. This system has the capability of developing schedules, station reports and string graphs with extreme accuracy. It is in use on Amtrak's Northeast Corridor, where some segments handle over 300 trains each weekday. The illustrations on the following pages are a computer generated string graph for the Union Pacific in Oregon and Idaho.

HISTORY DATA - GRANGER TO PORTLAND - 21 JAN 1996, SUN



TTMS - STRING GRAPH AND STICK MAP

HISTORY DATA - GRANGER TO PORTLAND - 21 JAN 1996, SUN



Connectivity:

Amtrak's entire system is a grid with connections generating a significant portion of the corporation's revenue. Some trains may have a significant number of passengers connecting to another train. There are many connecting points, but Chicago is the real hub of the Amtrak system. Sometimes an adjustment in one schedule will require a similar adjustment in a connecting schedule.

Negotiations:

Many people do not realize that Amtrak's schedules are negotiated with the railroads over which Amtrak trains operate. As noted at the beginning of this paper when Amtrak inherited the passenger train schedules in the early 1970's, some schedules were good, some schedules were bad and some overall were okay but needed restringing. Starting in 1976, Amtrak made a concerted effort to bring reasonable uniformity to the schedules in use throughout the system, but this is an ongoing process that demands time and personnel to stay in tune with operational changes.

Railroads are sensitive to schedules for two reasons: (1) they do not want passenger trains to interfere with their primary freight operations, and (2) on-time performance may drive millions of dollars in incentive payments. Both the length of schedules and schedule times are negotiated.

Railroads are sensitive to schedules for two reasons: (1) they do not want passenger trains to interfere with their primary freight operations, and (2) on-time performance may drive millions of dollars in incentive payments. Both the length of schedules and schedule times are negotiated.

Another critical factor is scheduling of commuter trains and intercity passenger trains so that conflicts are minimal. This is called slotting trains. Commuter trains and intercity trains may operate on three-minute headways, and passing locations are critical between express trains and local trains making all stops. In like manner, meeting opposing commuter trains is also a critical element.

The negotiation of schedules is as important as all other aspects of scheduling.

PART II

THE IMPACT OF DELAYS

The Impact of Delays

After consideration of all the factors essential to proper schedule construction, why don't trains run on time? The missing element is delay. Some delays are covered by recovery time; some are not. The causes of delays could fall into hundreds of categories, which in Amtrak's performance reporting system have generally been condensed into ten, plus "miscellaneous," which is a catch-all.

**When the amount of delay exceeds the available recovery time,
the train is late.**

Some delays are considered to be under Amtrak's control, such as:

- Mechanical failures
- Servicing delays
- Passenger related delays, the preponderance of which are excess station dwell times. This is a category of delay that must be carefully monitored by field supervision due to the tendency of dwell time delays increasing over a period of time.

Some delays are considered under the freight railroads' control:

- Freight train interference
- Signal delays
- Maintenance-of-way/slow order delays

Some delays are beyond the control of either organization:

- Firehoses over the tracks
- Snowslides, rockslides
- Heat orders, cold orders
- Grade crossing or trespasser accidents
- Vandalism
- 1,000 other things

These types of delays are known generally as "third-party" delays.

Some of the more bizarre delays Amtrak has experienced include striking and killing a buffalo in North Dakota, a World War II mine washing up on the shore of Puget Sound, and extracting an obese passenger from the toilet of the Broadway Limited.

Freight railroad causal factors generally account for 40 to 50 percent of the total delays -- while Amtrak-related delays generally account for 20 to 30 percent and all other delays for the remaining 20 to 30 percent

Signal Delays:

The frequency of signal delays closely follows freight train interference and is often the result of the signals working as intended because of a train ahead, as opposed to a true signal system failure.

It would be desirable to divide signal related delays into two sub-categories, i.e., working as intended or a signal failure. However accurate information is frequently not available until later and it would be difficult to segregate this information. When a signal is at "stop" because there is a freight train ahead, it should not be reported as a signal failure, it is freight train interference. Unfortunately, a crew delay report reading "Ajax, 10 minutes, red signal" will be reported by the person entering the data into the system as simply a signal delay for lack of any better information.

Slow Order Delays:

In like manner, maintenance-of-way/slow order delays are very significant delay factors and fall into two major categories:

The first is slow orders because of deteriorating track.

The second is slow orders or work orders because of routine maintenance or rehabilitation of track. Modern maintenance practices with production gangs frequently require taking the track out of service, or require slow orders and lookout orders which impact upon performance. However, just as in the case of signal delays, it would be desirable to sub-divide this category between slow orders due to track conditions and slow orders due to maintenance-of-way work. This should be easier to quantify and segregate than signal delays. However, train crews, and thus their delay reports, may not know or specify the causal factor of any given slow order.

Seasonal Schedules:

The issue of maintenance of way delays raises the question of whether train schedules should be modified to provide summer schedules (as opposed to winter schedules) to compensate for delays expected to result from major maintenance projects. The answer to this is clearly "yes" if the maintenance project will be of

sufficient duration to have a significant impact upon on-time performance. As new maintenance procedures have evolved with the use of large production gangs that take a segment of track out of service, and with the difficulties and production loss associated with clearing to allow a train or trains to pass, it is important to coordinate train operations and maintenance programs to minimize delays and to maximize productivity.

A more fundamental question is whether circumstances justify summer and winter schedules for more than just maintenance-of-way work; for example, ski resort travel that exists only in winter or other seasonal requirements, such as variance in station dwell time requirements.

Coordinating Maintenance Projects With Train Schedules:

Maintenance operations usually have a twelve-month window in the southern half of the United States, while the same work must be compressed into an engineering season, generally April to November, in the northern half of the States. Timetables generally change in early April and late October, which provides a nearly seven-month summer schedule and a slightly over five-month winter schedule. If a major 30-day maintenance-of-way project is going to occur, it is probably not practical to amend a schedule for seven months because of one month of impact. However, if a 90-day project is scheduled, a temporary schedule for that train, say June 1 to August 30, is

preferable to operating the train late every day or four or five days per week until the project is completed.

Coordinating maintenance-of-way projects with train operations is a relatively new concept that has not reached its full potential. Adjusting maintenance-of-way starting times or working night schedules may materially reduce potential delays; however it is equally important once this has been done to keep trains on time to protect or honor that window to minimize any maintenance-of-way delays as well as delays to the workforces.

There is frequently resistance to operating around or bussing around a major maintenance-of-way program. However it is often better to disrupt operations for a short period of time and let the work be promptly completed than to prolong the work and the delays resulting from a lengthy program.

Engineering Practices, Schedules, and Delays:

Since maintenance-of-way/slow order delays are a major factor in on-time performance, over the years there have been a number of poor practices that actually increase the amount of delay unnecessarily and which, if monitored judiciously, could significantly reduce delays.

Minimizing Delays Resulting From Maintenance-of-Way Practices:

1. Point restrictions (i.e., road crossing at MP 508.0 or bridge at MP 475.5) should be just that; however, we frequently find restrictions one mile either side of the point which results in a two mile restriction.
2. Frequently T&S restrictions will be imposed for several miles or for the entire length of the project even though the track has not been disturbed. Keeping the limits as short as possible and progressing as the work is done will minimize delays.
3. When slow orders result in a lower class of track, some carriers fail to recognize the passenger train speed differential (i.e., Class I: 10 mph freight, 15 mph passenger; Class II: 25 mph freight, 30 mph passenger).
4. Low speed rock and roll restrictions are not applicable to passenger equipment. Therefore, a 15 or 20 mph speed limit prohibition should not apply to passenger trains.
5. Failure to promptly lift restrictions once the condition is corrected unnecessarily delays trains. This is the result of either the M of W employee not notifying the dispatcher, or the dispatcher not promptly annulling the restriction.

6. Failure to properly place M of W flags will adversely affect the operation, as will a failure to properly remove the flags when work is completed. Frequently, the failure to place flags for slow orders has resulted in excessive delays.

7. Where speed restrictions apply while passing M of W gangs, the foreman may let the restriction apply throughout the entire work limits even though the gang has been passed or is working in one reasonably compact area.

8. When a significant or lengthy restriction applies on one main track, dispatchers may overlook crossing the trains over to an adjacent main track in traffic control territory to avoid the restriction and minimize delay when conditions permit. Dispatchers need to be much more sensitive to these delays.

9. In many cases there is a reluctance to go to 3-inch unbalanced curve speeds as it allows no tolerance. A waiver authorizing 4-inches will provide tolerance to go to a full 3-inches.

10. A series of curves may not be satisfactory for 35 mph in accordance with the FRA track safety standards, so they are restricted to 30 mph whereas an odd speed such as 33 mph would minimize delay. On one carrier, increasing the speed to 27 or 28 mph saved nearly 20 minutes. The slower the speed the more significant the time savings. Using odd speeds (other than ending with 5 or 0) can result in improved performance.

significant the time savings. Using odd speeds (other than ending with 5 or 0) can result in improved performance.

11. Heat orders are necessary when the ambient temperature exceeds certain predetermined limits. However, issuing heat orders the day before (i.e., 10:00 a.m. until 8:00 p.m.) even though the ambient temperature during those times is not known or equally as bad when the dispatcher fails to annul the restriction once the condition no longer exists, can result in significant delays. In one extreme case, a cold front with a cold rain moved through the area and even though the crew notified the dispatcher, nothing was done to remove the heat order.

All of the above points are the result of some actual experiences which have been encountered during Amtrak's riding programs. While not all of the above delays have occurred on all carriers, those railroads which actively police their properties to avoid such occurrences have not only benefited in minimizing delays to passenger trains and improving their incentive payments, but have also materially reduced delays to their freight operations. Unless a program is maintained to constantly monitor the operations, delays of the types listed above will occur.

There is actually a third category of maintenance-of-way delay that does not generally get recorded as a delay. Examples of this category include a long-standing slow order that eventually gets moved to a more permanent bulletin and, when no one is looking, slips from a bulletin into the new timetable and becomes a permanent restriction. A similar example is skimming some superelevation out of a series of curves and then through the same process slowing down the railroad -- a few miles at first, then a few more, until a hundred or more miles have been slowed down. Once the reduced speeds are entered into the timetable they are no longer reported as delay; however the overall impact is that they erode recovery time, until there are sufficient speed reductions which may result in trains running late. It can also result in thousands of minutes of delay each year that are not reported as delays. This practice is usually in violation of a carrier's contractual obligation to Amtrak, but restoration of speeds or other speed increases to compensate for the reduced speeds is difficult at best. This practice is analogous to taking a few cents out of the till each day, then a dollar or two, that nobody misses until an audit is performed. Delays are seldom reported in this area until a train cannot make its schedule. Sometimes this will be reported as a "running time" delay.

Passenger Train Meets:

The next major delay category is passenger trains or what is referred to as passenger train interference. This is the delay that occurs to one or both trains when two passenger trains meet. Freight railroads charge these delays against Amtrak since,

but for the presence of a second passenger train, the delay would not have occurred.

Amtrak charges these delays against the railroad, since train dispatching is performed by the railroad and time is allowed in the schedules to compensate for the meets.

While the responsibility for the delay may never be resolved, this becomes a stand-alone delay category over which reasonable control exists through good dispatching practices. On some lines where several pairs of passenger trains operate on single track, such as the San Diego Line or the San Joaquin Valley line, passenger train meets may be a leading cause of delay. Under these circumstances, one late train may cause the "domino effect," delaying the opposing trains that it is scheduled to meet. They in turn delay the opposing trains they are scheduled to meet. All this results from being at the wrong place at the wrong time. This raises the issue of "out of slot."

Out of Slot- A Myth or For Real?

There is a long-standing theory within the industry of the impact of a train being out of slot. The theory is that a train that is on-time is easier to keep on-time than it is for a late train to avoid becoming later.

It appears to be a fairly universal theory, or at least a perception, that once a train becomes late it is more difficult to keep on time than when it is on schedule.

However, even if this is true on dark (unsignalled) rail lines or perhaps in train order ABS (track warrant) territory, the very nature of a traffic control system (TCS) would mitigate against this being true in TCS territory.

For example, being late in TCS territory may simply move the meet from one controlled siding to another without material impact. It should be kept in mind that the degree to which scheduled freight trains actually adhere to their schedules is also a significant variable. Passenger train meets at locations other than the programmed siding may or may not have a negative impact upon one or both trains.

There are some out of slot factors that do have a direct impact upon a train getting later without regard to TCS territory -- for example, commuter trains where a late train may take multiple hits, or maintenance-of-way work or programmed work that would have been avoided if the train had been on time. In like manner, opposing or following fleets of freight trains could have an impact. Conversely, a late train might also avoid such delays by virtue of the fact that it is late, so the impact of being late cuts both ways.

As a result of these discussions, a study was conducted on whether the out of slot theory has any statistical merit in today's operation.¹

¹ See paper entitled "Out of Slot -- Fact or Myth," dated January 2, 1996.

The study addressed the performance of nearly 600 trains over seven railroads and totaled in excess of 500,000 train miles. The current Amtrak route system mileage includes approximately 3 percent which is unsignalled, 21 percent with an Automatic Block Signal System, and 76 percent with a Traffic Control System. Only 4 percent of the route mileage is over rail lines with both commuter and intercity passenger trains.

The study recognized that "out of slot" may be a problem in certain limited applications such as in a joint commuter/intercity operation.

The end result of that study was a conclusion that in today's operation the out of slot theory, except in limited applications such as stated above, is generally a myth. The percentage of on-time trains or late trains losing more time was not materially different, while late trains lost slightly less time than on-time trains. Whether the theory was ever more than a perception -- as opposed to having statistical validity -- cannot be proven since records no longer exist to substantiate the theory. One unexpected by-product of the study showed that trains are generally not abandoned from a dispatching standpoint once they become late.

Heat Orders and Cold Orders:

The use of heat orders and cold orders has grown significantly over recent years. The practice has considerable merit from an engineering standpoint, as it reduces stresses to the track structure during periods of extreme temperatures -- generally

ambient temperatures above 95 degrees F or below zero degrees F. Some carriers place heat or cold orders when there are extreme temperature changes within a prescribed period of time. Inconsistency is where the difficulty occurs. Generally, there is no standard practice, throughout the industry, either on when to apply such restrictions, or on the severity of such restrictions. Some carriers watch ambient temperatures carefully, place restrictions only when the temperature threshold is exceeded, and promptly remove restrictions when they no longer apply. Other carriers are not as precise. One carrier will issue a bulletin at midnight restricting train speeds from 10:00 a.m. until 8:00 p.m. the following day or days. Other carriers are responsive to placing the restrictions, but lax in removing them. On one carrier a long distance train was restricted by a heat order. A cold front passed through and the crew reported a cold rain to the dispatcher, who refused to annul the restriction, and the train received an unnecessary hour of delay.

Another issue is how such restrictions should be categorized. There are arguments that say they should be charged as maintenance-of-way delays, or weather related, or miscellaneous. Currently, they are generally captured as weather-related delays.

Avoidable Amtrak-Related Delays:

Another type of delay related to dwell time, but which often occurs under the cover of darkness and is probably underreported, is what is referred to as the "Grey

Poupon Syndrome." This occurs when two passenger trains meet on the road between stations and stop to hand off provisions to each other. These may include sheets or pillowcases, marinara sauce or Grey Poupon mustard and an occasional carryby passenger. This is not an accepted practice, but when it does occur it should be reported as a delay.

Similarly, the practice of "streetcarring" passengers, i.e., collecting tickets one at a time on the platform, which results in delay to the train, is not an acceptable practice and is probably never reported as a delay, but appears as excess dwell time and finally as a "passenger related" delay. This is a difficult delay to quantify -- i.e., what would the delay have been if normal boarding practices had occurred? Generally, it amounts to the number of minutes in excess of scheduled dwell time.

Killing Time:

Finally comes the question of "killing time" -- Is it a delay? The answer is "yes" it is a delay, but it is a good delay when used in moderation because it means the train is on time. Killing time comes in two forms. The first is running somewhat slower to avoid arriving at the next station too far ahead of scheduled departure time. The second is sitting at a station when the work is completed, waiting for the scheduled time to depart.

In any riding program it is essential to capture this time and record it as a delay. Killing one or two minutes is not evil; however when five or ten or more minutes are

frequently killed, it is a sign that the schedule needs to be restrung and the time used elsewhere where it will be more beneficial.

Late Connections:

This discussion addresses whether late connections are really a delay or whether they represent a duplicate counting of an accumulation of other delays.

Late connections are a delay en route if a train is held for a connecting train or bus. For example, Trains 8 and 28 connect at Spokane and Trains 92 and 82 are joined at Jacksonville, as are Trains 49 and 449 at Rensselaer. If a train is held for the other train, it is late-connection delay. Another example is in the San Joaquin Valley of California, where numerous busses connect with each train. Again, if a train is held at a point for a bus connection, it is a delay.

Conversely, if Carrier A delivers a train late each day to Carrier B at Denver, it is not a "delay" that occurs on Carrier B. It is a result of delays which actually occurred and are already accounted for on Carrier A. Another example: if a train is held in Chicago for connecting trains, it is late but it is not a delay that occurred on the carrier that received the train late.

The question may be raised, "So what? What difference does it make?" The answer lies in the example shown above at Denver. Carrier B reported thousands of

minutes of delay on late connections -- so many thousands of minutes, which did not occur on their railroad, that the statistics of what really happened on Carrier B were obscured by delays that actually occurred somewhere else. Carrier B used this as a subterfuge to obscure their own shortcomings. This type of obfuscation eventually led Amtrak to conduct the "out of slot" study described earlier, which showed that a late connection did not make a material difference in the train's performance. In the case of Carrier B, the late train did impact upon maintenance-of-way forces the next day but not on the overall performance of the train.

Just as some carriers may abuse or overreport late-connection delays, some conductors may enter "late arrival" as a delay on their delay reports just because the train arrives late at an intermediate crew-change point. However, the late arrival was caused by delays that already occurred on previous crew districts, and as such should have been accounted for by those previous crews.

Selective Delay Reporting:

Selective delay reporting is an old railroad practice of reporting only enough delay to cover the amount of time lost on schedule. This was generally done in the dispatcher's office. Does anyone recall seeing a delay reading "Ajax, 20 minutes, bad meet"? Not really. Perhaps "Ajax, 20 minutes, inspection," or "hard throwing switch" or "signal would not clear," but not a "bad meet."

Since all trains have recovery time, a train is not late until the amount of delay exceeds available recovery time. For example, if a train has 20 minutes available recovery time over a segment, it would generally (but not always), need 40 minutes of delay to be 20 minutes late. Selective delay reporting gave the dispatcher the opportunity to "select" which delays they would report to cover the 20 minutes of time lost on schedule.

The primary disadvantage of selective delay reporting was that only a portion of the delay was reported, so when it came to analyzing or identifying the problem areas the delay reports were worthless. Amtrak believes that all carriers have dropped this practice. It is essential to know where all the delays occur so that the problem areas that need to be addressed can be identified.

Today the Conductor's Delay Report is the document that Amtrak and most carriers recognize as the "official" delay report. Amtrak and one of the carriers conducted a detailed study in 1994 to evaluate delay reporting. The Conductor's Reports and dispatching reports were each compared to the actual delays as observed in the field by the study participants. It was concluded that the Conductor's Report was by far the more reliable document.

Train Dispatching:

The key to the performance of the trains rests with the train dispatcher.

Sometimes the dispatcher is faced with several different choices, such as:

Which train gets preference?
• No. 199 with the its intermodal trailers and containers • LABRF with its hot manifest • Madame X -- the division cleanup train • The Night Crawler - a puller handling empty camp cars • The Dogpatch Local - with two empty cars for the mill • Yard engine - going to the beanery to eat or switching the car repair tracks • All of the above
OR
• Passenger trains

What it all boils down to is the quality of train dispatching. The train dispatcher generally responds to the policies established by senior management. Which train gets the priority is the railroad's call.

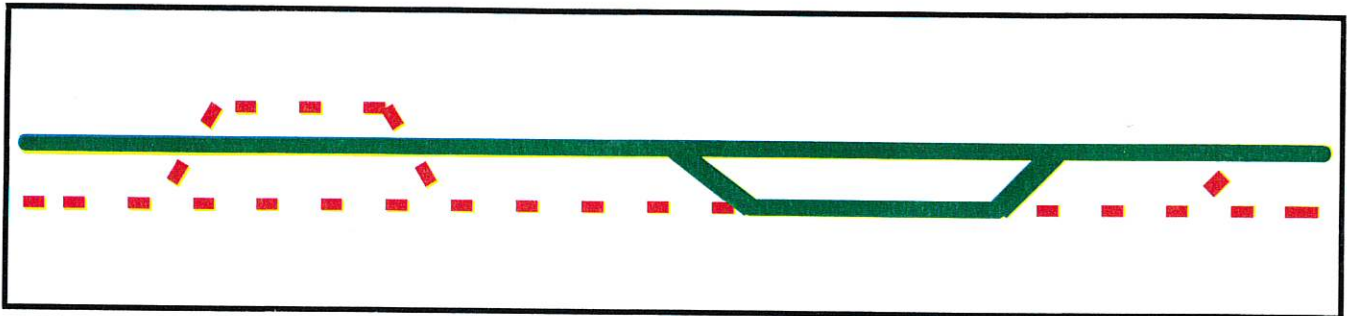
The policy varies significantly from one carrier to another. One key element is that passenger trains have priority over freight trains by statute, under the Rail

Passenger Service Act. All railroads profess to give passenger trains priority, but it just doesn't always happen. Statutory preference does not mean that passenger trains must always hold the main track at meeting points. On the contrary, in traffic control territory, the passenger train taking siding for a fast freight may expedite the meet and minimize the delay to both trains.

Computer Aided Dispatching or CAD (which perhaps should actually be known as Dispatcher Aided Computers) permits the computer to dispatch trains -- just plug in the priorities and watch them roll. Well, sort of. The computer can make some bad decisions unless the dispatcher catches the error and overrides the decision. Some dispatcher's territories are analogous to the old television skit of the one-man cake factory where cakes spew off a conveyor belt until chaos reigns. Conversely, where train dispatching territories have not been consolidated and expanded to the point where only the computer can dispatch the overall territory, CAD can be a valuable aid.

Software experts continually refine programs to correct situations where trains take siding with no train to meet -- known as meeting ghost trains -- or the local freight receives priority over the hottest train on the division. One thing generally agreed upon by all is that CAD-generated delay reports are for the most part worthless. Conversely, dispatchers can override the system and sometimes set up a situation that may negatively impact priority trains. It is a delicate balance.

Rationalization:



Rationalization, a buzzword of the 1980's and 90's, means reducing the physical plant to the maximum extent possible, often without regard to potential increases in traffic. It comes in two forms. The first is where mergers result in generally parallel routes. One is downgraded and the other becomes the primary line with resulting traffic increases. The second, which is actually more common, is where a conversion is made from two main tracks to single track. Sidings or main tracks are removed and the resulting plant has less capacity.

The situation is made worse if the segments of single track are made too long, if sidings or leftover segments of double track are made too short to enable "running meets" wherein neither opposing train has to stop for the other, or if sidings are left unsignalled or slow speed turnouts are installed, requiring trains to crawl into the sidings at slow speed. All of these "short-cuts" may lessen the initial cost of the single-tracking project or increase short-term savings, but by creating a less-fluid main line, they lower overall capacity and increase long-term costs and delays.

The increase in traffic over the past several years has left many railroad managements scratching their heads about how to increase capacity quickly and economically. Some carriers with critical choke points are now hurting in trying to expedite freight traffic in addition to passenger trains. Threading a passenger train through congested areas is a difficult task.

Summary:

Finally, there is a logical question: If the train runs late every day, why not add time to fix it? There are three basic reasons:

- 1) If there is excess time in the schedule the train dispatchers tend to take liberties with the trains.
- 2) Station work tends to expand to fill all available time.
- 3) If a schedule is basically adequate and a train runs 30 minutes late each day and 30 minutes is simply added to the schedule, in 90 days it will be running 30 minutes late each day. It is essential to fix the basic problem.

The opening comments stated that a train with too much time in the schedule will run late. The reason for this is that the first dispatcher will put the passenger train in the siding and advance a freight train, thus causing, for example, 25 minutes delay, and

using the rationale that there is an abundance of time in the schedule to make up the delay. The second dispatcher's rationale is "I got the train 25 minutes late and I delivered it 25 minutes late. I broke even with no delay to report." The third dispatcher says: "Damn, every day I get the train late and they expect me to make it up." And so the train runs late. In many instances, implementing faster schedules has actually improved performance.

A well strung schedule with a reasonable amount of recovery time is the first step to improve performance. The second step is receiving the appropriate priority and a commitment from the carrier to maintain good on-time performance. The third step is good train dispatching.

A sufficient amount of recovery time is very important. Several years ago, the schedule of a long-distance train on one carrier was tightened up so much that almost all recovery time was removed. The dispatchers discovered that no matter how hard they tried they could not reliably keep the train on time. Soon they gave up and began giving priority to freight trains. After the real problem of insufficient recovery time was identified and after recovery time was added back into the schedule, dispatching and on-time performance both improved dramatically.

But excess time in the schedule can be a performance killer. One on carrier, 30 minutes was added to the schedules of four trains to compensate for a major engineering project. On-time performance improved substantially. However, once the

project was completed, performance dropped as the excess time was used for flagrant freight train interference which far exceeded the additional 30 minutes. On another carrier, it was pointed out that delays were occurring because the two passenger trains met on a single track and that with only a one-hour adjustment, the meet could be moved to double track. This was done, but the performance did not improve. Passenger train interference dropped significantly, but freight train interference increased dramatically.

Fortunately, the examples in the preceding paragraph are atypical for the industry. There is a concerted effort being made between carriers and Amtrak to improve on-time performance, to identify and evaluate operating problems, and to better control the delays that are controllable.

Scheduling Trains: A Summary

Scheduling has three basic components and an important secondary factor:

- Pure running time
- Dwell time
- Recovery time
- Passenger train meets

Once the above data are known:

- The schedule is strung.
- It is checked on a string graph.

Connections are a critical part of scheduling

Once proposed schedules are developed they are negotiated with the freight railroads:

- Schedule times
- Length of schedules

When delay time exceeds recovery time, the train is late.

Delays are measured in broad categories:

- Amtrak related delays
- Freight carrier related delays
- Passenger train meets
- Delays beyond control of either party

Control is vested in the train dispatcher, who governs priorities consistent with corporate policy. The dispatcher has the ultimate control.

Schedules must be coordinated with both freight and commuter operations.

Schedules impact upon the carriers' ability to earn millions of dollars in incentives.